

Audio-Visual Camera Pose Estimation with Passive Scene Sounds and In-the-Wild Video

Daniel Adebi¹, Sagnik Majumder¹, and Kristen Grauman¹

The University of Texas at Austin, Austin TX 78712, USA
{ikadebi,sagnik,grauman}@cs.utexas.edu

Abstract. Understanding camera motion is a fundamental problem in embodied perception and 3D scene understanding. While visual methods have advanced rapidly, they often struggle under visually degraded conditions such as motion blur or occlusions. In this work, we show that passive scene sounds provide cues complementary to vision for relative camera pose estimation for in-the-wild videos. We introduce a simple but effective audio-visual framework that integrates direction-of-arrival (DOA) spectra and binauralized embeddings into a state-of-the-art vision-only pose estimation model. Our results on two large datasets show consistent gains over strong visual baselines, plus robustness when the visual information is corrupted. To our knowledge, this represents the first work to successfully leverage audio for relative camera pose estimation in real-world videos, and it establishes incidental, everyday audio as an unexpected but promising signal for a classic spatial challenge.

Keywords: Camera pose estimation · Audio-visual learning · Multi-modal learning

1 Introduction

Consider a video captured at a dark, crowded concert where someone wearing a camera moves through the venue searching for friends near the stage. The visual feed is severely degraded by dim lighting, motion blur, and constantly shifting crowds, making it difficult to track camera movement through vision alone. Yet the surrounding audio tells a clear spatial story: music grows louder as the person approaches the stage and the music shifts as they turn their head; conversations fade in and out with proximity; and ambient crowd noise provides persistent directional cues. We refer to such footage—video and audio captured under unconstrained, everyday conditions, without staging or simulation, and subject to natural visual and acoustic variation—as *in-the-wild*.

Could a *camera pose estimation model* use such incidental sounds in real-world video to refine its spatial story too? This goal is appealing for multiple reasons. First, incidental *passive* scene sounds are compelling because—unlike actively emitted chirps for echolocation [12, 19, 70]—they are non-intrusive, varied, and freely available with video. Second, everyday sounds are invariant to some of the key obstacles for traditional vision-only camera pose estimation, like



Fig. 1: Main idea: we propose to estimate relative camera pose from in-the-wild videos using both vision and multichannel audio. Unlike traditional active echolocation sensing, our model relies only on passive scene audio—gaining spatial cues opportunistically from naturally occurring ambient and foreground sound sources.

poor lighting, motion blur, textureless regions, and occlusions. Audio’s geometric cues remain reliable even when vision fails.

Motivated by these observations, we introduce an audio-visual camera pose estimation model. Given two video frames and their respective multi-channel (meaning at least two channels) audio excerpts, the objective is to predict the relative camera pose separating the two frames. See Fig. 1. Building on a state-of-the-art vision-only model [16], we introduce a joint audio representation comprised of explicit direction-of-arrival audio cues and a mono-to-binaural “lifted” embedding. Importantly, we successfully train entirely with real-world video from a variety of scenarios and activities—a significant departure from the status quo, where methods exploring spatial audio rely on static (usually 3D-scanned [51,58]) environments and simulated audio [10, 11, 18, 19, 50, 70].

We validate our approach on two datasets. Our model achieves significant improvements, outperforming the state-of-the-art [11] and establishes (to our knowledge) the first-ever results for audio-visual camera pose estimation on challenging real-world videos. Our key contributions are:

- We introduce a spatial audio encoder that learns spatial embeddings from two complementary cues: direction-of-arrival spectra and cross-view audio binauralization features.
- We validate our approach on the large-scale Ego-Exo4D dataset [23] containing hundreds of subjects and environments and the HM3D-SS [10, 51] dataset, achieving state-of-the-art performance among audio-visual baselines for relative camera pose estimation.
- We provide a comprehensive analysis of how different audio characteristics affect multimodal camera pose estimation.

Overall, our work establishes incidental, passive audio as a promising signal for camera pose.

2 Related Work

2.1 Vision-Based Camera Pose Estimation

Accurate relative camera pose estimation is fundamental to a wide range of applications in 3D vision, augmented reality, and robotics. Most camera pose estimation methods rely solely on visual information [5, 6, 15, 16, 28, 32, 36, 52, 54, 64–67, 74, 75], and limited prior work incorporates depth [15, 75], IMU [7], or optical flow [67]. Earlier methods such as SuperGlue [54] and LoFTR [59] rely on correspondence through keypoint matching, which makes them well suited to scenarios with high degrees of visual overlap between images. More recent approaches like DUST3R [65] and Reloc3r [16] leverage end-to-end transformer architectures, allowing direct correspondence estimation without relying on explicit keypoint detection and matching. FAR [52] integrates feature matching with learning-based approaches to improve performance under low visual overlap scenarios.

All of these methods rely solely on visual input for camera pose prediction, leveraging training data that spans simulated environments [29, 49, 55, 58, 61], reconstructed 3D scans of real scenes [4, 8, 13, 72, 73], photographs of static real-world environments [3, 31, 35, 41, 57, 62], and real-world videos [14, 33, 60, 69, 76]. In our work, we incorporate audio as an additional modality to support pose estimation.

2.2 Echolocation For Camera Poses

To our knowledge, the only prior work to leverage audio for camera pose estimation does so in a very different setting—by actively emitting sounds into the environment in order to collect echoes [70]—and is limited to training and testing in simulation. Furthermore, no prior work uses passive scene audio for relative camera pose estimation. Our method relies solely on naturally occurring sounds already present in the scene, without emitting signals, generating new sounds in a simulated scene, or requiring controlled acoustic conditions, making it passive, unobtrusive, and compatible with real-world videos.

2.3 Localizing Sound Sources

Leveraging both sound and vision to localize sound sources has emerged as a powerful strategy [11, 30, 39, 68]. Recent work shows how camera motion can reinforce estimates of sound source locations [11, 39]. Results are done in simulation [11] or using real-world video egomotion to supervise binaural sound localization [39]. In related spatial problems, audio-visual models can learn floorplan maps [50] and 3D scene structure [12] by associating sounding objects with likely rooms and/or listening for echolocation responses. Other work trains navigation policies to move to a sounding object in an unmapped environment [10, 18]. Unlike any of the above, our aim is to compute accurate 6 DoF relative camera poses, not sound source locations or environment maps.

2.4 Audio-visual Self-Supervised Pretraining

Extensive prior work investigates how vision and audio interact to learn self-supervised representations, using the two modalities to create pretext tasks and improve training [1, 2, 26, 43–47]. These studies build tasks around synthesis [46, 47], cross-modal alignment [1, 2, 19, 26, 45], and masked auto-encoding approaches [21, 22, 25], typically targeting semantic downstream problems such as audio-visual event classification and retrieval. Different from these objectives, our approach incorporates spatial audio-based representations into camera pose estimation.

3 Task Motivation and Definition

The spatial sound perceived in a 3D scene is shaped by the relative location and orientation of the microphones used in capturing the sound, relative to the sound source(s), as well as the scene’s geometry and materials. Consequently, any translation and/or rotation of the microphones results in changes in the sound’s spatial attributes, which reveal the nature and extent of microphone motion. Based on this knowledge, we hypothesize that leveraging spatial audio **in addition to vision** for relative camera pose prediction in in-the-wild videos—where the audio is captured using two or more microphones that move rigidly with the camera—can improve the pose prediction quality over using vision alone. We anticipate audio to be particularly valuable when vision alone is unreliable due to factors like occlusion, bad lighting, or camera failure, which are common in real-world video. Meanwhile, a sound source that is not visible but is still audible adds useful cues—e.g., someone speaking or a radio located out of view.

To validate our hypothesis, we propose a novel audio-visual in-the-wild relative camera pose estimation task. In this task, we consider a video clip $V = (I, A)$, where I and A denote the visual and multi-channel audio streams, respectively. The visual stream I consists of N RGB image frames, such that $I = \{I_1, \dots, I_N\}$. The audio stream A is time-synchronous with I and consists of N audio segments, such that $A = \{A_1, \dots, A_N\}$. Each audio segment A_j has C audio channels, such that $A_j = \{A_j^1, \dots, A_j^C\}$, and is extracted by sampling a fixed-length time window (1000ms) centered at frame I_j , making A_j temporally aligned with V_j . Note that the shorter the time window, the less likely there are substantial changes in sound source positions intra-window.

We make no assumptions about the audio other than it being captured from two (or more) microphones co-located with the camera.¹ Such audio-visual capture is increasingly available in not only research devices like Aria [17] glasses but also commercial consumer devices like GoPro, Insta360 Go, and RayBan Meta glasses. In particular, we are interested in leveraging *passive* sounds—whatever are the incidental sounds in the scene from various sources, foreground or background—as opposed to *active* sounds emitted into the environment to

¹ We use the terms spatial audio and multi-channel audio interchangeably, to refer to audio captured from two or more microphones.

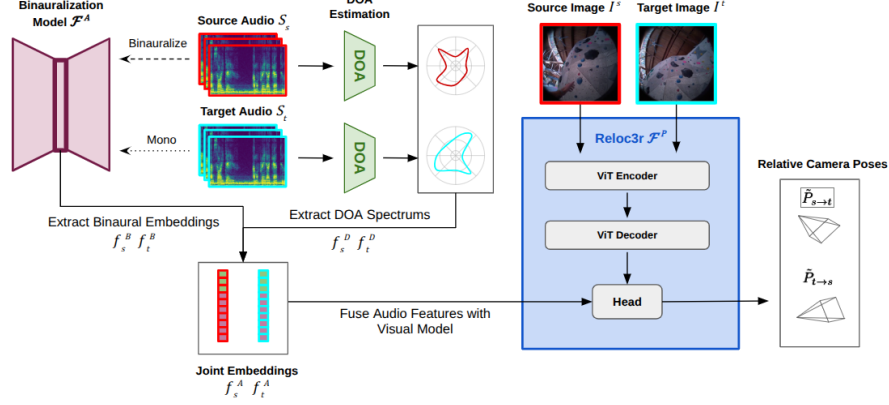


Fig. 2: We extend the Reloc3r [16] architecture with our Spatial Audio Encoder (SAE). Given a pair of source and target frames, coupled with their corresponding synchronized audio clips, our network predicts relative camera poses for the input image pair, in both directions (from source to target, and from target to source). The proposed SAE produces a combination of analytical and learned spatial audio embeddings that are subsequently integrated with Reloc3r’s visual features via late fusion, enabling the network to leverage complementary audio-visual cues for performing high-quality pose prediction. This design allows joint learning of translation and rotation and additionally provides robustness to visual corruption, as we show in results.

perform sensing, i.e., for echolocation. The advantage of passive sounds is that they are non-intrusive and “free”, and available even on previously recorded video over which our models have no control.

Given a pair of source and target RGB frames (I_s, I_t) and their corresponding multi-channel audio segments (A_s, A_t), the goal in this task is to learn a model $\mathcal{F}$, such that $\mathcal{F}(I_s, A_s, I_t, A_t) = \hat{P}_{s \rightarrow t}$, where $\hat{P}_{s \rightarrow t}$ is an estimate of the 6-DoF relative camera pose $P_{s \rightarrow t}$ between I_s and I_t . The matrix $P \in \mathbb{R}^{3 \times 4}$ represents the rotation and translation of the pose change. Formulating relative camera pose between two views is standard [11, 16, 52, 70]. It also serves as a core primitive underlying recent multi-view and temporally-extended systems [5, 64, 65], and is natural for online scenarios where future multi-view context is unavailable.

4 Approach

Our method is composed of two modules: **1)** a spatial audio encoder $\mathcal{F}^A$, and **2)** an audio-visual relative camera pose predictor $\mathcal{F}^P$. See Fig. 2. While the audio encoder $\mathcal{F}^A$ is responsible for extracting rich spatial cues from the multi-channel audio inputs by using both analytical and learned representations, the pose predictor $\mathcal{F}^P$ combines these spatial audio cues with visual cues using a SOTA camera pose prediction backbone [16] to estimate the relative camera pose. Next, we describe the design of these modules and their training.

4.1 Spatial audio encoder

Given any multi-channel audio segment A_j (cf. Sec. 3), our spatial audio encoder $\mathcal{F}^A$ first computes a short-time Fourier transform (STFT) and represents the audio as a spectrogram S_j . S_j is a matrix, such that $S_j \in \mathbb{R}^{C \times F \times W}$, where F and W are the number of frequency bins and time windows, respectively, and C is the number of audio channels (e.g., $C = 2$ for stereo or binaural audio).

We compute two representations from the spectrogram: the direction-of-arrival (DOA) and a learned spatial audio embedding. For the former, we use an analytical direction-of-arrival (DOA) estimator [53] to compute $f_j^D = \mathcal{F}^D(S_j)$, such that $f_j^D \in \mathbb{R}^{360}$. This DOA spectrum explicitly captures the distribution of sound energy around the microphone location, revealing the direction(s) of major sound sources in the environment with respect to the camera. For the latter, we compute a binaural audio embedding $f_j^B = \mathcal{F}^B(S_j)$ using a learned extractor (defined below) to provide an implicit but strong representation of the audio’s spatial characteristics.

Finally, the encoder fuses f_j^D and f_j^B into a spatial audio embedding f_j^A . The joint embeddings for the audio from both the source and target poses are passed to the camera pose predictor $\mathcal{F}^P$ for further processing. Whereas the DOA feature f_j^D signals the direction of major sound sources in the surroundings, the binaural audio features capture the full fabric of all overlapping sounds and their spatial layout. Importantly, even overlapping sound sources and ambient noises are a *signal* for our camera pose estimation problem—not a nuisance as in traditional tasks like audio source separation. The synergy of the two complementary features, enabled by their fusion, is essential for high-quality pose prediction, as we show in results. Next, we elaborate on the design of these encoders and our feature fusion strategy.

Direction-of-arrival spectrum. For our direction-of-arrival (DOA) feature, we use the MUSIC [56] algorithm with frequency normalization [53], henceforth referred to as MUSIC++. This method has been shown to produce reliable DOA estimates in the presence of multiple sound sources. Given a multi-channel audio-spectrogram S_j , MUSIC++ analytically computes a one-dimensional DOA feature f_j^D , such that $f_j^D \in [0, 1]^{360}$ and each entry measures the average audio intensity across the corresponding 1° -wide sector in the azimuth circle. In essence, f_j^D captures the angular variation in the audio intensity relative to the direction in which the microphones are facing. Notably, this DOA distribution will shift predictably as a function of the camera’s motion—and vice versa—provided there are not sudden changes in the scene sounds between the time frame I_s and I_t are captured (1 sec. in our experiments).

Learned binaural audio features. Inspired by prior work for self-supervised audio(-visual) learning [9, 11, 19, 20, 38, 42, 48, 71], we train an embedding to lift monaural to binaural sound in the target viewpoint. The idea is to surface essential geometric cues from the audio. Specifically, we train $\mathcal{F}^B$ by attempting to solve a novel viewpoint² (view) acoustic synthesis (NVAS) task, where given a

² Here, “viewpoint” refers to a microphone pose.

pair of views—one source and one target—the goal is to learn a model that can take the binaural audio for the source view and the monaural audio for the target view as inputs, and generate the binaural audio for the target view as output. By training the NVAS model on a large and challenging dataset [23] comprising a large variety of dynamic 3D scenes with multiple overlapping sound sources at diverse locations and ambient noise, as well as arbitrary viewpoint changes, $\mathcal{F}^B$ learns to encode rich spatial audio features from two-channel audio.

Spatial audio feature fusion. To leverage the complementary benefits offered by these two features, we fuse f_j^D and f_j^B into a single spatial audio feature f_j^A through concatenation, such that $f_j^A = [f_j^D || f_j^B]$. Our fusion strategy, coupled with the use of non-linearities in the subsequent model layers (see Sec. 4.2), allows for full feature mixing, making the fused features highly expressive and conducive to high-quality camera pose estimation.

4.2 Audio-visual relative camera pose predictor

For predicting the relative camera pose estimate $\tilde{P}$, we adapt a state-of-the-art vision-only relative camera pose predictor, Reloc3r [16], such that it can take both visual and spatial audio inputs. Given a pair of RGB images (I_s, I_t) , the original Reloc3r model processes them using a pair of visual branches with the exact same architecture and shared weights, in order to predict their relative camera poses in both directions, $\tilde{P}_{s \rightarrow t}$ and $\tilde{P}_{t \rightarrow s}$. Towards that goal, it feeds each image to just one branch, which extracts visual features from the image using an encoder-decoder pair. Additionally, it uses cross-attention between the decoders of the two branches to establish fine-grained spatial correspondences between the corresponding visual features. It then uses a stack of linear layers as its camera pose prediction head, which leverages the aforementioned correspondences and estimates the camera pose of each input image relative to the other.

To adapt Reloc3r for audio-visual inputs, we feed an audio-visual feature f^{AV} to the camera pose estimation head of each branch, where f^{AV} is produced by concatenating the head’s original input feature f^V and our spatial audio feature f^A , such that $f^{AV} = [f^A || f^V]$. Doing this late fusion of the visual and spatial audio features allows our pose predictor to learn important spatial correlations between the visual and audio modalities. Furthermore, it preserves the visual backbone’s cues: informative audio can augment vision, but uninformative audio has limited ability to override strong visual features.

4.3 Model training

Binaural audio feature extractor training. Following [11, 20], given a target viewpoint $\mathcal{T}$ and its ground-truth binaural audio segment $A_{\mathcal{T}}^B$, our NVAS model, whose encoder we use for our binaural feature extractor $\mathcal{F}^B$, predicts an estimate $\hat{S}_{\mathcal{T}}^A$ of the spectrogram of the ground-truth difference between the left and right channels of the target audio, $S_{\mathcal{T}}^A$, such that $S_{\mathcal{T}}^A$ is computed by performing STFT on $A_{\mathcal{T}}^l - A_{\mathcal{T}}^r$. With this formulation, the model is not required to predict the

exact value of each entry in spectrogram, which greatly speeds up training while still enabling the model to learn strong binaural features due to its requirement of accurately predicting any inter-channel differences. We train the NVAS model and consequently, our binaural feature extractor F^B , by computing $\mathcal{L}^B$, where $\mathcal{L}^B$ is the L_1 loss between $S_{\mathcal{T}}^{\Delta}$ and $\tilde{S}_{\mathcal{T}}^{\Delta}$, such that

$$\mathcal{L}^B = \|\tilde{S}_{\mathcal{T}}^{\Delta} - S_{\mathcal{T}}^{\Delta}\|_1. \quad (1)$$

Camera pose estimator training. Given an RGB image pair (I_s, I_t) and the corresponding relative camera pose ground-truths $(P_{s \rightarrow t}, P_{t \rightarrow s})$, we train our audio-visual camera pose predictor F^P using $\mathcal{L}^P$, where $\mathcal{L}^P$ is the mean of the pose prediction loss $L_{s \rightarrow t}^P$ for $P_{t \rightarrow s}$ and its converse $L_{t \rightarrow t}^P$, corresponding to the prediction loss for $P_{t \rightarrow s}$, such that

$$\mathcal{L}^P = L_{s \rightarrow t}^P + L_{t \rightarrow s}^P. \quad (2)$$

Following [16], our pose prediction loss L^P has two components, measuring the error in the rotation and translation predictions, respectively.

5 Experiments

We present experiments validating our model on three datasets, (1) the in-the-wild video dataset Ego-Exo4D [23], (2) the perceptually realistic HM3D-SoundSpaces [10, 51], which allows comparison against the SOTA audio-visual method, and (3) test sets with added visual corruptions for both.

5.1 Experimental Setup

Datasets. To assess our method’s performance on real-world data, we first train and evaluate our models on the **Ego-Exo4D dataset** [23], using all 132 hours of egocentric sequences for which there are audio recordings (Ego-Exo4D’s exo cameras are static and hence not relevant for pose estimation). Audio in Ego-Exo4D is captured with 7 microphones on the Aria glasses, which provide synchronized multichannel audio and precise camera pose annotations. The data spans a diverse range of scenarios, motions, and activities, including indoor navigation, social interactions, and object manipulation, enabling evaluation across varied acoustic conditions like speech vs. object sounds, and foreground sounds vs. ambient noise, and visual conditions such as occlusions and lighting changes. We note that this complex video content pushes the boundary compared to how relative camera pose is typically assessed in simulated settings and/or with static scenes. See Sec. 7 in supplemental for more detailed analysis on the camera pose and audio distribution statistics.

To generate our dataset, we uniformly sample frame pairs that are one second apart, ensuring sufficient motion while maintaining temporal continuity. This sampling choice offers a consistent supervision signal without restricting generality, as the model is trained across a wide variety of activities, environments, and

motion magnitudes. We also evaluate our model performance on different temporal splits in Tab. 7 (Supp.) to showcase how the model generalizes to different time gaps between input frames. We follow the official Ego-Exo4D data splits, ensuring that frame pairs used for training and testing are sampled from entirely disjoint video subsets, fully preventing cross-split overlap during evaluation.

Secondly, to enable direct comparison with SLfM [11], which generates camera orientation and sound source locations, we additionally train and evaluate a variant of our approach on the Habitat-Matterport3D SoundSpaces (**HM3D-SS**) dataset [11, 51], which combines Habitat-Matterport3D’s photorealistic indoor scenes with the SoundSpaces [10] framework in order to produce realistic spatial audio in simulation. This variant uses 2 microphones instead of 7. We follow the same data splits and scene configurations used by SLfM [11].

Finally, to assess robustness to corrupted visual inputs, we also generate a **corrupted test set** for each of the above datasets by applying a random combination of photometric and geometric transformations to each frame, where both the chosen transformations and their strengths are randomly decided. The set of possible perturbations includes Gaussian blur, Gaussian noise, and color jitter (brightness, contrast, saturation, and hue adjustments). See Tab. 3 for examples. These degradations include both standard corruption types used in prior robustness work [24, 34, 37, 40] as well as more extreme perturbations that occur when the transformations are applied with full strength to simulate unexpected visual failures, enabling evaluation under both typical and worst-case visual conditions.

Baselines. We compare against the following baselines and state-of-the-art methods:

- **Chance:** this is a heuristic that simply outputs the mean of the ground-truth relative camera poses from the validation set
- **Naive Audio Camera Pose (NA-CP):** this baseline predicts the relative camera pose by replacing the RGB inputs to a Reloc3r [16] model with corresponding multi-channel audio spectrograms, and finetunes it to predict estimates of the corresponding ground-truth camera poses.
- **Direction-of-arrival Camera Pose (DOA-CP):** this baseline first computes the direction-of-arrival (DOA) spectrum, and then predicts the relative camera pose by processing the DOA spectrum using a stack of linear layers.
- **Reloc3r [16]:** a state-of-the-art vision-only relative camera pose predictor that takes pairs of RGB images as inputs and uses a transformer encoder-decoder model to predict the camera pose for the image pair, in both directions.
- **Reloc3r [16]-AV:** an audio-visual version of Reloc3r that performs late fusion (cf. Sec. 4.2) of Reloc3r’s visual features and the audio features from the corresponding layer of our NA-CP baseline, and feeds the fused features to Reloc3r’s pose prediction head to estimate the relative camera pose. We finetune this model using pairs of audio-visual inputs and relative camera pose ground-truths.
- **Reloc3r [16]-AV Cross Attention:** we leverage multimodal cross-attention to integrate audio and visual features, allowing the model to selectively

Table 1: Comparison of relative camera pose estimation performance on Ego-Exo4D dataset validation split. We report AUC@5, AUC@10, and AUC@20 metrics for rotation only, translation only, and total based on the maximum of the rotation and translation error. Best performing results are in bold, and second best results are underlined. Relative gains are based on performance over vision-only model. All gains over baselines and ablations are statistically significant ($p \leq 0.05$).

Methods	Rotation Only			Translation Only			Total		
	AUC@5	AUC@10	AUC@20	AUC@5	AUC@10	AUC@20	AUC@5	AUC@10	AUC@20
<i>Vision Only</i>									
Reloc3r [16]	37.35	53.62	68.98	0.73	2.63	7.77	0.57	2.33	7.33
<i>Audio Only</i>									
Chance	<u>19.02</u>	<u>33.33</u>	50.03	0.06	0.26	1.11	0.02	0.12	<u>0.76</u>
NA-CP Model	18.83	33.27	<u>50.05</u>	<u>0.07</u>	<u>0.29</u>	<u>1.16</u>	<u>0.02</u>	<u>0.12</u>	0.72
DOA CP Model	19.07	33.37	50.07	0.09	0.37	1.38	0.03	0.14	0.78
<i>Audio + Vision</i>									
Reloc3r [16]-AV	37.10	53.30	68.72	0.88	2.97	8.53	0.69	2.65	8.07
Reloc3r [16]-AV Cross-Att	<u>38.01</u>	53.82	69.01	0.92	3.15	8.89	<u>0.76</u>	2.83	8.34
Ours w/ Monaural	37.98	53.80	68.98	0.94	3.18	8.89	0.74	2.82	8.37
Ours w/o DOA	37.48	53.45	68.88	<u>0.96</u>	<u>3.21</u>	<u>8.92</u>	0.75	<u>2.86</u>	<u>8.41</u>
Ours w/o Binaural	37.99	<u>53.87</u>	69.04	0.93	3.12	8.80	0.74	2.80	8.34
Ours	38.54	54.49	69.61	1.03	3.36	9.34	0.81	2.99	8.82
Relative Gains (%)	+3.2%	+1.6%	+0.9%	+41.1%	+27.8%	+20.2%	+42.1%	+28.3%	+20.3%
<i>Corrupted Vision</i>									
Reloc3r [16]	25.17	38.88	52.80	0.34	1.28	4.28	0.21	1.00	3.66
<i>Audio + Corrupted Vision</i>									
Reloc3r [16]-AV	26.17	39.76	53.59	0.43	1.57	4.96	0.29	1.27	4.36
Reloc3r [16]-AV Cross-Att	26.38	39.83	53.68	0.42	1.61	5.08	<u>0.32</u>	1.29	4.44
Ours w/ Monaural	<u>27.00</u>	<u>40.28</u>	<u>53.89</u>	0.44	1.61	5.07	0.31	1.28	4.42
Ours w/o DOA	26.14	39.76	53.71	<u>0.45</u>	1.61	5.06	0.31	1.29	4.42
Ours w/o Binaural	26.56	40.06	53.81	0.44	<u>1.62</u>	<u>5.10</u>	0.29	<u>1.30</u>	<u>4.48</u>
Ours	27.17	40.66	54.46	0.49	1.77	5.45	0.35	1.44	4.80
Relative Gains (%)	+8.0%	+4.4%	+3.1%	+44.1%	+38.3%	+27.3%	+66.7%	+44.0%	+31.2%

fuse information across modalities for improved scene understanding. This approach is inspired by recent work in audio-visual learning, where cross-attention enables effective alignment between visual cues and audio signals [27, 63].

- **Ours w/ Monaural:** an ablation of our approach where we modify our binauralization model so that it predicts single-channel audio. From here, we extract the monaural embeddings from our audio encoder to fuse with the rest of our model. DOA spectra are still used as an input to this model ablation.

Evaluation metrics. Following previous research [16, 54, 59, 66], we evaluate our model on the Ego-Exo4D dataset using AUC@5/10/20. These metrics compute the area under the accuracy curve (AUC) for camera pose estimation, with thresholds of 5, 10 and 20 degrees on the maximum angular error in the rotation and translation predictions. This provides a comprehensive measure of relative pose accuracy across varying error tolerance levels. For the HM3D-SS dataset, we strictly follow the evaluation protocol used in prior work [11] and report the mean absolute error (MAE) in predictions of the rotation azimuth.

Table 2: Results on the HM3D-SS dataset test split. We report rotation accuracy using mean absolute error (MAE ($^{\circ}$)) for the original data (left) and corrupted visual signals (right), following the evaluation protocol of SLfM [11]. Lower values are better. Relative performance gains are based on performance over vision-only model. All gains over baselines and ablations are statistically significant ($p \leq 0.05$).

Method	Original Corrupted	
Chance	29.41	—
SLfM [11]	0.77	—
NA-CP Model	2.56	
DOA-CP Model	28.63	
Reloc3r [16]	0.09	3.83
Reloc3r [16]-AV	0.09	3.28
Reloc3r [16]-AV Cross-Att	0.08	2.82
Ours w/ Monaural	0.08	2.89
Ours w/o DOA	0.08	2.78
Ours w/o Binaural	0.08	2.66
Ours	0.06	2.46
Relative Gains (%)	+33.3%	+35.8%

5.2 Results

We present the results for both datasets, including their visually corrupted variants, and then provide analysis of what audio properties are most and least amenable to augmenting vision for camera pose estimation.

Audio-visual relative camera pose estimation. Table 1 reports relative camera pose estimation results on the Ego-Exo4D test split.³ Among audio-only models, the DOA-CP model achieves the strongest overall performance, with higher AUC scores across most thresholds compared to both the mean baseline and the naive audio model. This demonstrates that direction-of-arrival (DOA) features extracted from multichannel raw audio provide meaningful geometric cues for estimating relative pose, even without visual supervision.

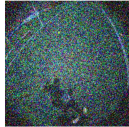
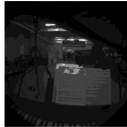
When combining audio information with visual inputs, our approach outperforms all of our baselines and ablations with statistical significance ($p \leq 0.05$), with significant relative gains in total AUC over all thresholds (42%/28%/20% for AUC@5/10/20 respectively). We can see that incorporating structured audio representations leads to consistent improvements across both rotation and translation metrics. Specifically, the addition of DOA features substantially enhances rotation accuracy, reflecting their ability to encode spatial geometry complementary to vision, while the binaural embeddings provide cues that improve translation estimates. Our approach using monaural embeddings worsens performance, showcasing the importance of inferred binauralization in our method; the graceful degradation highlights its practical utility for even single-microphone devices. Our full model achieves the highest total AUC across all thresholds, demonstrating the benefits of jointly leveraging both embeddings.

Furthermore, Tab. 1 shows that under corrupted vision conditions (bottom third of the table), our model maintains significantly higher performance than

³ Note that the SLfM [11] method is not applicable to Ego-Exo4D because when it jointly trains for sound source localization and camera motion, it requires ground truth direction-of-arrival values, which Ego-Exo4D does not provide.

Table 3: AUC@20 scores demonstrating robustness under different visual corruption scenarios, with corresponding visualizations below. Our method consistently outperforms the baseline for all forms of corruption with statistical significance ($p < 0.05$).

σ Noise	Vision Only	Ours	σ Blur	Vision Only	Ours	Jit. Strength	Vision Only	Ours
0.00	7.33	8.82	0	7.33	8.82	0.0	7.33	8.82
0.05	6.29	7.62	1	5.82	7.61	0.2	7.27	8.79
0.10	4.82	5.92	2	5.07	6.21	0.4	6.55	7.97
0.20	4.28	5.55	4	3.52	4.41	0.6	6.32	7.67
0.30	4.08	5.52	8	2.63	3.47	0.8	6.20	7.64

$\sigma = 0$	$\sigma = 0.3$	$\sigma = 0$	$\sigma = 8$	JS = 0	JS = 0.8
					
(a) Gaussian Noise		(b) Gaussian Blur		(c) Color Jitter	

all baselines, highlighting its robustness when visual cues are degraded. These results collectively confirm that our multi-embedding fusion not only improves overall accuracy but also enhances cross-modal reliability in challenging visual environments. Table 3 breaks down the relative accuracy for our method and its vision-only counterpart for multiple levels of different categories of corruption. Note that we consistently achieve higher scores than the vision-only baseline across all forms of degradation at all corruption levels.

Having analyzed the impact of visual corruption, we next examine the impact of audio corruption at inference time. Unlike in source separation, our method does not depend on clean or intelligible audio content; rather, it gains signal from any inter-channel differences, so additional sound energy, even if perceptually messy, still offers usable spatial cues. Indeed, we find that with noisy audio (Gaussian noise), performance remains largely preserved (Total AUC@5/10/20 = 0.78/2.88/8.62), well above the vision-only baseline (0.57/2.33/7.33; Tab. 1). This is again consistent with our late-fusion design (Sec. 4.2), which limits the influence any single corrupted modality can have on the final prediction.

Figure 3 provides some qualitative examples comparing our model to the vision-only baseline. Here we see the impact that ambient audio has across many scenarios. Our method can adapt to sudden changes in pose that seem to throw off our vision-only model. There are a few cases where our method struggles, though, specifically when a single dominant audio starts up (or stops happening) between frames.

Table 2 shows the results on the HM3D-SS dataset, for both the original data (left) and the visually corrupted (right). We can see our model achieves the lowest mean absolute rotation error, outperforming both the state-of-the-art SLfM [11] and the Reloc3r [16] baseline in both normal and corrupted visual settings. While the improvement over the visual-only baseline is minimal, this is largely due to the controlled nature of HM3D-SS scenes, where the visual signal alone is highly informative. This reinforces our emphasis on real-world video as a more challenging and essential testbed.

We also observe that the DOA Camera Pose model performs poorly in this setting, which may stem from inaccuracies in how direction-of-arrival features

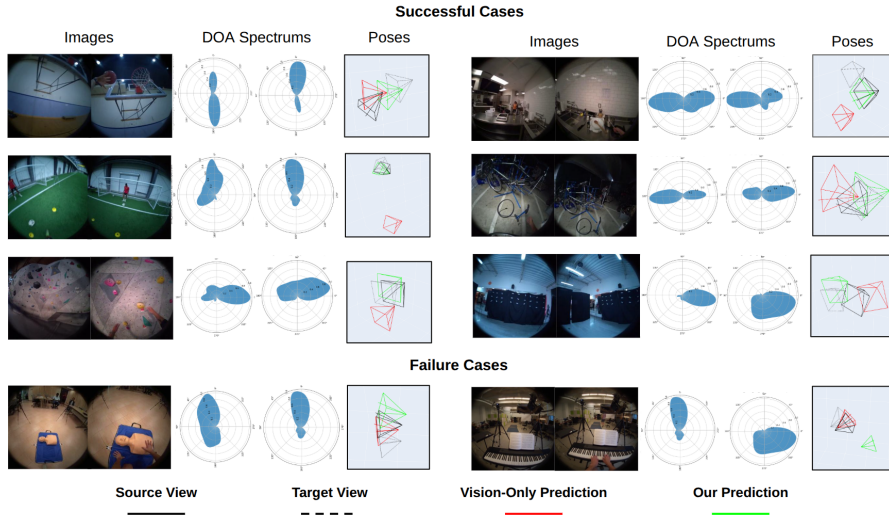


Fig. 3: Qualitative examples of our audio-visual method outperforming a vision-only state-of-the-art relative camera pose estimation model. Our model performs best when there are consistent audio signals in the scene that aid in localization, such as the soccer video in the second row. Some failure cases of our model include when a single dominant audio starts up in one frame but does not exist in a previous frame (e.g., bottom right, where the musician doesn’t start playing music til the second frame).

are computed in simulation and the absence of precise microphone position parameters in the SoundSpaces [10] configuration. Nonetheless, our DOA-enhanced fusion model consistently surpasses the naive audio variant and vision-only baseline, suggesting that DOA cues remain beneficial even under imperfect calibration and strong visual dominance.

Decoupling camera motion and audio change. Table 4a evaluates performance in regimes where camera motion and acoustic change are explicitly decoupled, and in heavy translation settings. To construct these subsets, we automatically filter sequences using camera pose statistics, defining large motion as translation over 1m or rotation greater than 20° (above the mean of the dataset). We separately quantified sound variation to detect large acoustic changes (e.g. a sound exists in one frame but does not exist in second frame). We then cross-selected clips satisfying one condition while constraining the other, and performed a light manual verification pass to ensure clean separation between motion- and audio-dominant cases.

Across these controlled settings, our method consistently outperforms Vision-Only. Gains are particularly evident when sound sources change substantially, even under large camera motion, indicating robustness to acoustic variability. Improvements also persist when camera motion exceeds 1m translation but remains under 10° rotation, demonstrating strong performance in high-translation,

Table 4: Analysis of multimodal performance under controlled motion/audio regimes (left) and across diverse Ego-Exo4D scenarios (right).**(a)** AUC@5/10/20 under varying camera motion and acoustic change settings.

Method	AUC@5	AUC@10	AUC@20
<i>Large Translation >1m, Small Rotation <10°</i>			
Vision-Only	2.56	7.58	16.65
Ours	3.57	9.51	20.08
<i>Small Camera Motion, Large Sound Change</i>			
Vision-Only	0.71	2.95	9.15
Ours	1.05	3.74	10.89
<i>Large Camera Motion, Small Sound Change</i>			
Vision-Only	0.30	1.36	4.99
Ours	0.40	1.57	5.89

(b) Impact of scene audio characteristics across scenarios. The middle column lists each scenario’s predominant audio characteristic(s), identified via manual inspection. The rightmost column reports the percentage of frames where our method outperforms the vision-only (Reloc3r) baseline.

Scenario	Audio type(s)	Audio helps? (%)
Soccer	★ ■ ▲	84.63
Music	◆ ● ▲	81.40
Bike Repair	◆ ■ ▲	79.55
Rock Climbing	★ ◆ ● ▲	79.41
Basketball	★ ◆ ▲	72.15
Health	◆ ■	67.16
Dance	◆ ▲	66.14
Cooking	● ■	64.13
Audio Types: ★ Far-field ◆ Near-field ● Dominant Single ■ Low Signal ▲ Frequent Changes		

low-rotation regimes where directional audio cues provide especially informative constraints on relative displacement.

Analysis of passive in-the-wild scene audio characteristics. Table 4b summarizes representative audio characteristics observed across the Ego-Exo4D scenarios and illustrates that the proposed approach consistently improves performance across multiple real-world acoustic settings represented in the dataset. For each scenario, we annotate the predominant audio characteristics most frequently observed in the data, including far-field source audio (distant sound sources), near-field or egocentric audio (sounds proximate to the camera), dominant single sound sources, low or minimal audio signals, and frequent changes in sound content between consecutive frames. We then measure the frequency at which our method outperforms the vision-only state-of-the-art baseline (“Audio Helps? (%)” column). For more details on how we annotate our audio characteristics, see Tab. 9 in the supplemental.

Our gains are strongest in scenes with multiple spatially distributed sound sources, where DOA spectra and binauralized embeddings provide complementary geometric cues to vision. Scenarios with the most concurrent sources yield our most frequent gains (80-85%, Tab. 9 in Supp). Even under omnidirectional sound (flat DOA spectra), our binaural embeddings still allow gains over vision only by 8%/6%/3% for AUC@5/10/20. In Soccer and Music, concurrent sources (e.g., ball impacts, footsteps, instruments, people talking) generate distinct directional signatures that evolve with camera motion, enabling our audio-visual model to better resolve relative pose—especially under visual degradation. Bike Repair similarly benefits from spatially separated tool interactions. In contrast, Cooking and Dance often contain fewer independently localized emitters, more centralized acoustic activity, or a single dominant sound source, reducing directional diversity, though consistent sound cues still yield meaningful improvements over vision-only baselines. Further analysis can be found in Sec. 7 in the supplemental.

Does audio truly provide a geometric signal? To investigate whether the audio branch exploits geometrically meaningful cues rather than scene-, activity-, or dataset-specific correlations, we perform a series of negative-control experiments that progressively disrupt the correspondence between audio and camera motion. Performance degrades as the audio becomes increasingly inconsistent with the underlying scene geometry: synchronized audio achieves the highest AUC@5 (0.81), followed by temporally offset audio (1–5s; 0.68), same-scene shuffled audio (0.66), and cross-scene audio replacement (0.50), which performs worse than the vision-only baseline (0.57). These results indicate that dataset priors and semantic correlations are not enough—our results depend critically on the geometric consistency between the audio observations and the visual scene.

Failure modes. While our audio-visual framework demonstrates strong robustness across diverse scenarios, it remains susceptible to specific failure modes, particularly when acoustic cues become disjointed from the underlying scene geometry. Our model can struggle in near-silent or diffuse-source conditions where directional cues are inherently weak. Furthermore, because the framework relies on the geometric consistency of the audio, artificially disrupting this alignment, e.g. inserting an unrelated soundtrack in post-production, would actively mislead the model and degrade performance. Abrupt acoustic changes that occur independently of camera motion also present a challenge (see Fig. 3, bottom row).

6 Conclusion

We present a multimodal approach for relative camera pose estimation in video that exploits spatial cues from multichannel audio. Our method achieves state-of-the-art performance among audio-visual baselines on both real-world and simulated settings, demonstrating that ambient audio contains rich spatial structure to complement visual signals for a classic spatial vision task.

While our approach is effective, it has some limitations that present opportunities for future work. Translation estimation remains challenging in visually degraded or acoustically complex scenarios. In addition, it would be interesting to extend the model to explicitly gate the audio’s influence depending on its properties (e.g., stationarity). Exploring cross-dataset and cross-device generalization is also an important area for future work. Despite these challenges, our work establishes that audio-visual fusion substantially improves pose estimation across diverse acoustic conditions. We hope this work motivates further research in leveraging sound for geometric reasoning and robust scene understanding.

Acknowledgements

This research is supported in part by Lyda Hill Philanthropies. We thank the UT Austin Vision Group, Ziyang Chen (U. Michigan), and Kazuki Shimada (Sony) for helpful discussions.

References

1. Afouras, T., Owens, A., Chung, J.S., Zisserman, A.: Self-supervised learning of audio-visual objects from video. In: Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part XVIII 16. pp. 208–224. Springer (2020)
2. Arandjelović, R., Zisserman, A.: Look, listen and learn. 2017 IEEE International Conference on Computer Vision (ICCV) pp. 609–617 (2017)
3. Arnold, E., Wynn, J., Vicente, S., Garcia-Hernando, G., Monszpart, Á., Prisacariu, V.A., Turmukhambetov, D., Brachmann, E.: Map-free visual relocalization: Metric pose relative to a single image. In: ECCV (2022)
4. Baruch, G., Chen, Z., Dehghan, A., Dimry, T., Feigin, Y., Fu, P., Gebauer, T., Joffe, B., Kurz, D., Schwartz, A., Shulman, E.: Arkitscenes - a diverse real-world dataset for 3d indoor scene understanding using mobile rgb-d data. In: NeurIPS (2021), <https://arxiv.org/pdf/2111.08897.pdf>
5. Cabon, Y., Stoffl, L., Antsfeld, L., Csurka, G., Chidlovskii, B., Revaud, J., Leroy, V.: Must3r: Multi-view network for stereo 3d reconstruction. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 1050–1060 (2025)
6. Cai, R., Zhang, J.Y., Henzler, P., Li, Z., Snavely, N., Martin-Brualla, R.: Can generative video models help pose estimation? In: CVPR (2025)
7. Cerezo, S., Civera, J.: Camera motion estimation from RGB-D-inertial scene flow. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW) (2024), <https://arxiv.org/abs/2404.17251>
8. Chang, A., Dai, A., Funkhouser, T., Halber, M., Niessner, M., Savva, M., Song, S., Zeng, A., Zhang, Y.: Matterport3d: Learning from rgb-d data in indoor environments. International Conference on 3D Vision (3DV) (2017)
9. Chen, C., Richard, A., Shapovalov, R., Ithapu, V.K., Neverova, N., Grauman, K., Vedaldi, A.: Novel-view acoustic synthesis. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 6409–6419 (2023)
10. Chen, C., Schissler, C., Garg, S., Kobernik, P., Clegg, A., Calamia, P., Batra, D., Robinson, P.W., Grauman, K.: Soundspaces 2.0: A simulation platform for visual-acoustic learning. In: NeurIPS 2022 Datasets and Benchmarks Track (2022)
11. Chen, Z., Qian, S., Owens, A.: Sound localization from motion: Jointly learning sound direction and camera rotation. In: International Conference on Computer Vision (ICCV) (2023), <https://arxiv.org/abs/2303.11329>
12. Christensen, J.H., Hornauer, S., Yu, S.X.: Batvision: Learning to see 3d spatial layout with two ears. 2020 IEEE International Conference on Robotics and Automation (ICRA) pp. 1581–1587 (2019), <https://api.semanticscholar.org/CorpusID:209376315>
13. Dai, A., Chang, A.X., Savva, M., Halber, M., Funkhouser, T., Nießner, M.: Scannet: Richly-annotated 3d reconstructions of indoor scenes. In: Proc. Computer Vision and Pattern Recognition (CVPR), IEEE (2017)
14. Damen, D., Doughty, H., Farinella, G.M., Fidler, S., Furnari, A., Kazakos, E., Moltisanti, D., Munro, J., Perrett, T., Price, W., et al.: The epic-kitchens dataset: Collection, challenges and baselines. IEEE Transactions on Pattern Analysis and Machine Intelligence **43**(11), 4125–4141 (2020)
15. Ding, Y., Kocur, V., Vávra, V., Haladová, Z.B., Yang, J., Sattler, T., Kukełova, Z.: Reposed: Efficient relative pose estimation with known depth information. arXiv preprint arXiv:2501.07742 (2025)

16. Dong, S., Wang, S., Liu, S., Cai, L., Fan, Q., Kannala, J., Yang, Y.: Reloc3r: Large-scale training of relative camera pose regression for generalizable, fast, and accurate visual localization. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 16739–16752 (June 2025)
17. Engel, J., Somasundaram, K., Goesele, M., Sun, A., Gamino, A., Turner, A., Tatlatof, A., Yuan, A., Souti, B., Meredith, B., et al.: Project aria: A new tool for egocentric multi-modal ai research. *arXiv preprint arXiv:2308.13561* (2023)
18. Gan, C., Zhang, Y., Wu, J., Gong, B., Tenenbaum, J.B.: Look, listen, and act: Towards audio-visual embodied navigation. In: *ICRA* (2020)
19. Gao, R., Chen, C., Al-Halab, Z., Schissler, C., Grauman, K.: Visualechoes: Spatial image representation learning through echolocation. In: *ECCV* (2020)
20. Gao, R., Grauman, K.: 2.5d visual sound. In: *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 324–333 (2019). <https://doi.org/10.1109/CVPR.2019.00041>
21. Georgescu, M.I., Fonseca, E., Ionescu, R.T., Lucic, M., Schmid, C., Arnab, A.: Audiovisual masked autoencoders. *arXiv preprint arXiv:2212.05922* (2022)
22. Gong, Y., Rouditchenko, A., Liu, A.H., Harwath, D., Karlinsky, L., Kuehne, H., Glass, J.R.: Contrastive audio-visual masked autoencoder. In: *The Eleventh International Conference on Learning Representations* (2023), <https://openreview.net/forum?id=QPtMRyk5rb>
23. Grauman, K., Westbury, A., Torresani, L., Kitani, K., Malik, J., Afouras, T., Ashutosh, K., Baiyya, V., Bansal, S., Boote, B., et al.: Ego-exo4d: Understanding skilled human activity from first-and third-person perspectives. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 19383–19400 (2024)
24. Hendrycks, D., Dietterich, T.: Benchmarking neural network robustness to common corruptions and perturbations. *arXiv preprint arXiv:1903.12261* (2019)
25. Huang, P.Y., Sharma, V., Xu, H., Ryali, C., Fan, H., Li, Y., Li, S.W., Ghosh, G., Malik, J., Feichtenhofer, C.: Mavil: Masked audio-video learners. *arXiv preprint arXiv:2212.08071* (2022)
26. Korbar, B., Tran, D., Torresani, L.: Cooperative learning of audio and video models from self-supervised synchronization. *Advances in Neural Information Processing Systems* **31** (2018)
27. Lee, J., Chang, J., Lee, D., Choi, J.: Ca²st: Cross-attention in audio, space, and time for holistic video recognition. *arXiv preprint arXiv:2503.23447* (2025)
28. Leroy, V., Cabon, Y., Revaud, J.: Grounding image matching in 3d with mast3r. In: *European conference on computer vision*. pp. 71–91. Springer (2024)
29. Li, W., Saeedi, S., McCormac, J., Clark, R., Tzoumanikas, D., Ye, Q., Huang, Y., Tang, R., Leutenegger, S.: Interiornet: Mega-scale multi-sensor photo-realistic indoor scenes dataset. In: *British Machine Vision Conference (BMVC)* (2018)
30. Li, Z., Zhao, B., Yuan, Y.: Cyclic learning for binaural audio generation and localization. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 26669–26678 (June 2024)
31. Li, Z., Snavely, N.: Megadepth: Learning single-view depth prediction from internet photos. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*. pp. 2041–2050 (2018)
32. Li, Z., Tucker, R., Cole, F., Wang, Q., Jin, L., Ye, V., Kanazawa, A., Holynski, A., Snavely, N.: Megasam: Accurate, fast and robust structure and motion from casual dynamic videos. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 10486–10496 (2025)

33. Ling, L., Sheng, Y., Tu, Z., Zhao, W., Xin, C., Wan, K., Yu, L., Guo, Q., Yu, Z., Lu, Y., et al.: D13dv-10k: A large-scale scene dataset for deep learning-based 3d vision. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 22160–22169 (2024)
34. Liu, J., Wang, Z., Ma, L., Fang, C., Bai, T., Zhang, X., Liu, J., Chen, Z.: Benchmarking object detection robustness against real-world corruptions. *International Journal of Computer Vision* **132**(10), 4398–4416 (2024)
35. Liu, X., Tayal, P., Wang, J., Zarzar, J., Monnier, T., Tertikas, K., Duan, J., Toisoul, A., Zhang, J.Y., Neverova, N., Vedaldi, A., Shapovalov, R., Novotny, D.: Uncommon objects in 3d. In: arXiv (2024)
36. Loiseau, T., Bourmaud, G., Lepetit, V.: Alligat0r: Pre-training through co-visibility segmentation for relative camera pose regression. CoRR **abs/2503.07561** (March 2025), <https://doi.org/10.48550/arXiv.2503.07561>
37. Ma, S., Zhang, J., Cao, Q., Tao, D.: Posebench: Benchmarking the robustness of pose estimation models under corruptions. ArXiv **abs/2406.14367** (2024), <https://api.semanticscholar.org/CorpusID:270619440>
38. Majumder, S., Al-Halah, Z., Grauman, K.: Learning spatial features from audio-visual correspondence in egocentric videos. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 27058–27068 (June 2024)
39. Min, A., Chen, Z., Zhao, H., Owens, A.: Supervising sound localization by in-the-wild egomotion. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 23936–23946 (2025)
40. Mintun, E., Kirillov, A., Xie, S.: On interaction between augmentations and corruptions in natural corruption robustness. In: Beygelzimer, A., Dauphin, Y., Liang, P., Vaughan, J.W. (eds.) *Advances in Neural Information Processing Systems* (2021), <https://openreview.net/forum?id=L0Hyqjfyra>
41. Mirowski, P., Banki-Horvath, A., Anderson, K., Teplyashin, D., Hermann, K.M., Malinowski, M., Grimes, M.K., Simonyan, K., Kavukcuoglu, K., Zisserman, A., Hadsell, R.: The streetlearn environment and dataset. CoRR **abs/1903.01292** (2019), <http://arxiv.org/abs/1903.01292>
42. Morgado, P., Li, Y., Vasconcelos, N.: Learning representations from audio-visual spatial alignment. In: NeurIPS (2020)
43. Morgado, P., Li, Y., Nvasconcelos, N.: Self-supervised generation of spatial audio for 360° video. *Advances in Neural Information Processing Systems* **33**, 4733–4744 (2020)
44. Ngiam, J., Khosla, A., Kim, M., Nam, J., Lee, H., Ng, A.: Multimodal deep learning. In: *International Conference on Machine Learning* (2011)
45. Owens, A., Efros, A.A.: Audio-visual scene analysis with self-supervised multisensory features. In: *Proceedings of the European Conference on Computer Vision (ECCV)*. pp. 631–648 (2018)
46. Owens, A., Isola, P., McDermott, J., Torralba, A., Adelson, E.H., Freeman, W.T.: Visually indicated sounds. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*. pp. 2405–2413 (2016)
47. Owens, A., Wu, J., McDermott, J.H., Freeman, W.T., Torralba, A.: Ambient sound provides supervision for visual learning. In: *European Conference on Computer Vision* (2016)
48. Parida, K.K., Srivastava, S., Sharma, G.: Beyond mono to binaural: Generating binaural audio from mono audio with depth and cross modal attention. In: *2022 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*. pp. 2151–2160 (2022). <https://doi.org/10.1109/WACV51458.2022.00221>

49. Puig, X., Undersander, E., Szot, A., Cote, M.D., Partsey, R., Yang, J., Desai, R., Clegg, A.W., Hlavac, M., Min, T., Gervet, T., Vondrus, V., Berges, V.P., Turner, J., Maksymets, O., Kira, Z., Kalakrishnan, M., Malik, J., Chaplot, D.S., Jain, U., Batra, D., Rai, A., Mottaghi, R.: Habitat 3.0: A co-habitat for humans, avatars and robots (2023)
50. Purushwalkam, S., Gari, S.V.A., Ithapu, V.K., Schissler, C., Robinson, P., Gupta, A.K., Grauman, K.: Audio-visual floorplan reconstruction. 2021 IEEE/CVF International Conference on Computer Vision (ICCV) pp. 1163–1172 (2020), <https://api.semanticscholar.org/CorpusID:229923178>
51. Ramakrishnan, S.K., Gokaslan, A., Wijmans, E., Maksymets, O., Clegg, A., Turner, J.M., Undersander, E., Galuba, W., Westbury, A., Chang, A.X., Savva, M., Zhao, Y., Batra, D.: Habitat-matterport 3d dataset (HM3d): 1000 large-scale 3d environments for embodied AI. In: Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (2021), <https://arxiv.org/abs/2109.08238>
52. Rockwell, C., Kulkarni, N., Jin, L., Park, J.J., Johnson, J., Fouhey, D.F.: Far: Flexible, accurate and robust 6dof relative camera pose estimation. In: CVPR (2024)
53. Salvati, D., Drioli, C., Foresti, G.L.: Incoherent frequency fusion for broadband steered response power algorithms in noisy environments. *IEEE Signal Processing Letters* **21**(5), 581–585 (2014)
54. Sarlin, P.E., DeTone, D., Malisiewicz, T., Rabinovich, A.: SuperGlue: Learning feature matching with graph neural networks. In: CVPR (2020), <https://arxiv.org/abs/1911.11763>
55. Savva, M., Kadian, A., Maksymets, O., Zhao, Y., Wijmans, E., Jain, B., Straub, J., Liu, J., Koltun, V., Malik, J., Parikh, D., Batra, D.: Habitat: A Platform for Embodied AI Research. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) (2019)
56. Schmidt, R.: Multiple emitter location and signal parameter estimation. *IEEE Transactions on Antennas and Propagation* **34**(3), 276–280 (1986). <https://doi.org/10.1109/TAP.1986.1143830>
57. Snavely, N., Seitz, S.M., Szeliski, R.: Photo tourism: Exploring photo collections in 3d. In: SIGGRAPH Conference Proceedings. pp. 835–846. ACM Press, New York, NY, USA (2006)
58. Straub, J., Whelan, T., Ma, L., Chen, Y., Wijmans, E., Green, S., Engel, J.J., Mur-Artal, R., Ren, C., Verma, S., Clarkson, A., Yan, M., Budge, B., Yan, Y., Pan, X., Yon, J., Zou, Y., Leon, K., Carter, N., Briales, J., Gillingham, T., Mueggler, E., Pesqueira, L., Savva, M., Batra, D., Strasdat, H.M., Nardi, R.D., Goesele, M., Lovegrove, S., Newcombe, R.: The Replica dataset: A digital replica of indoor spaces. arXiv preprint arXiv:1906.05797 (2019)
59. Sun, J., Shen, Z., Wang, Y., Bao, H., Zhou, X.: LoFTR: Detector-free local feature matching with transformers. CVPR (2021)
60. Sun, P., Kretschmar, H., Dotiwalla, X., Chouard, A., Patnaik, V., Tsui, P., Guo, J., Zhou, Y., Chai, Y., Caine, B., Vasudevan, V., Han, W., Ngiam, J., Zhao, H., Timofeev, A., Ettinger, S., Krivokon, M., Gao, A., Joshi, A., Zhang, Y., Shlens, J., Chen, Z., Anguelov, D.: Scalability in perception for autonomous driving: Waymo open dataset. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (June 2020)
61. Szot, A., Clegg, A., Undersander, E., Wijmans, E., Zhao, Y., Turner, J., Maestre, N., Mukadam, M., Chaplot, D., Maksymets, O., Gokaslan, A., Vondrus, V., Dharur,

- S., Meier, F., Galuba, W., Chang, A., Kira, Z., Koltun, V., Malik, J., Savva, M., Batra, D.: Habitat 2.0: Training home assistants to rearrange their habitat. In: *Advances in Neural Information Processing Systems (NeurIPS)* (2021)
62. Thomee, B., Shamma, D.A., Friedland, G., Elizalde, B., Ni, K., Poland, D., Borth, D., Li, L.J.: Yfcc100m: The new data in multimedia research. *Communications of the ACM* **59**(2), 64–73 (2016)
63. Venkatraman, S., Sharma, V., Malarvannan, S., Narendra, M., et al.: Multimodal emotion recognition using audio-video transformer fusion with cross attention. *arXiv preprint arXiv:2407.18552* (2024)
64. Wang, J., Chen, M., Karaev, N., Vedaldi, A., Rupprecht, C., Novotny, D.: Vggt: Visual geometry grounded transformer. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (2025)
65. Wang, S., Leroy, V., Cabon, Y., Chidlovskii, B., Revaud, J.: Dust3r: Geometric 3d vision made easy. In: *CVPR* (2024)
66. Wang, Y., He, X., Peng, S., Tan, D., Zhou, X.: Efficient LoFTR: Semi-dense local feature matching with sparse-like speed. In: *CVPR* (2024)
67. Wimbauer, F., Chen, W., Muhle, D., Rupprecht, C., Cremers, D.: Anycam: Learning to recover camera poses and intrinsics from casual videos. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 16717–16727 (2025)
68. Wu, X., Wu, Z., Ju, L., Wang, S.: Binaural audio-visual localization. In: *AAAI*. vol. 35, pp. 2961–2968 (2021)
69. Xu, N., Yang, L., Fan, Y., Yang, J., Yue, D., Liang, Y., Price, B., Cohen, S., Huang, T.: Youtube-vos: Sequence-to-sequence video object segmentation. In: *Proceedings of the European conference on computer vision (ECCV)*. pp. 585–601 (2018)
70. Yang, K., Firman, M., Brachmann, E., Godard, C.: Camera pose estimation and localization with active audio sensing. In: *European Conference on Computer Vision*. pp. 271–291. Springer (2022)
71. Yang, K., Russell, B., Salamon, J.: Telling left from right: Learning spatial correspondence of sight and sound. In: *CVPR* (2020)
72. Yao, Y., Luo, Z., Li, S., Zhang, J., Ren, Y., Zhou, L., Fang, T., Quan, L.: Blend-edmvs: A large-scale dataset for generalized multi-view stereo networks. *Computer Vision and Pattern Recognition (CVPR)* (2020)
73. Yeshwanth, C., Liu, Y.C., Nießner, M., Dai, A.: Scannet++: A high-fidelity dataset of 3d indoor scenes. In: *Proceedings of the International Conference on Computer Vision (ICCV)* (2023)
74. Yin, Z., Shi, J.: Geonet: Unsupervised learning of dense depth, optical flow and camera pose. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (June 2018)
75. Zhang, J., Herrmann, C., Hur, J., Jampani, V., Darrell, T., Cole, F., Sun, D., Yang, M.H.: MonST3r: A simple approach for estimating geometry in the presence of motion. In: *The Thirteenth International Conference on Learning Representations* (2025), <https://openreview.net/forum?id=1JpqxFgWCM>
76. Zhou, T., Tucker, R., Flynn, J., Fyfe, G., Snavely, N.: Stereo magnification: Learning view synthesis using multiplane images. *arXiv preprint arXiv:1805.09817* (2018)